

МІНІСТЕРСТВО ОСВІТИ І НАУКИ УКРАЇНИ
НАЦІОНАЛЬНА АКАДЕМІЯ НАУК УКРАЇНИ
НАЦІОНАЛЬНИЙ ТЕХНІЧНИЙ УНІВЕРСИТЕТ УКРАЇНИ «КИЇВСЬКИЙ
ПОЛІТЕХНІЧНИЙ ІНСТИТУТ ІМЕНІ ІГОРЯ СІКОРСЬКОГО»
ХАРКІВСЬКИЙ НАЦІОНАЛЬНИЙ УНІВЕРСИТЕТ РАДІОЕЛЕКТРОНІКИ
НАЦІОНАЛЬНИЙ ЦЕНТР «МАЛА АКАДЕМІЯ НАУК УКРАЇНИ»

НАУКОЄМНІ ТЕХНОЛОГІЇ ОПТИМІЗАЦІЇ ТА КЕРУВАННЯ В ІНФОКОМУНІКАЦІЙНИХ МЕРЕЖАХ

Колективна монографія

Під загальною редакцією
В.М. Безрука, Л.С. Глоби, О.Є. Стрижака

Київ

2019

УДК 621.391+004.7

НЗ4

Рекомендовано до друку Вченою радою Національного центру «Мала академія наук України» (протокол № 1 від 15 січня 2019 року), науково-технічною радою Харківського національного університету радіоелектроніки (протокол № 2 від 25.06.2018 р.)

Рецензенти:

Бараннік В.В., д.т.н., проф., начальник кафедри бойового застосування та експлуатації АСУ Харківського Національного університету Повітряних Сил імені Івана Кожедуба; Климаш М.М., д.т.н., проф., завідувач кафедри телекомунікацій Національного університету «Львівська політехніка».

НЗ4

Наукоємні технології оптимізації та керування в інфокомунікаційних мережах : монографія / Під загальною редакцією В.М. Безрука, Л.С. Глоби, О.Є. Стрижака. – К.: Інститут обдарованої дитини НАПН України, 2019. – 194 с.
ISBN 978-617-7734-02-3

В монографії викладені сучасні результати досліджень ряду авторів з наукоємних технологій оптимізації та управління в інфокомунікаційних мережах. Розглядаються особливості методів багатокритеріальної оптимізації систем з урахуванням сукупності показників якості, а також особливості та результати вирішення задач управління при передачі та обробці інформації для різних типів інфокомунікаційних мереж.

УДК 621.391+004.7

© Національний технічний університет України «Київський політехнічний інститут імені Ігоря Сікорського», 2019

© Харківський національний університет радіоелектроніки, 2019

© Національний центр «Мала академія наук України», 2019

ISBN 978-617-7734-02-3 © Інститут обдарованої дитини НАПН України, 2019

ЗМІСТ

	Вступ	4
1	МЕТОДИ БАГАТОКРИТЕРІАЛЬНОЇ ОПТИМІЗАЦІЇ МЕРЕЖ ЗВ'ЯЗКУ <i>Безрук В.М.</i>	7
2	БАГАТОКРИТЕРІАЛЬНИЙ ВИБІР ПРОЕКТНИХ ВАРІАНТІВ ПРИ ПЛАНУВАННІ СИСТЕМ МОБІЛЬНОГО ЗВ'ЯЗКУ <i>Безрук В.М., Чеботарьова Д.В.</i>	27
3	ВИБІР ПЕРЕВАЖНИХ ВАРІАНТІВ ЗАСОБІВ ЗВ'ЯЗКУ НА ОСНОВІ МЕТОДУ АНАЛІЗУ ІЄРАРХІЙ <i>Безрук В.М., Власова В.О., Скорик Ю.В.</i>	47
4	БАГАТОКРИТЕРІАЛЬНИЙ ПІДХІД ДО ОПТИМАЛЬНОЇ МАРШРУТИЗАЦІЇ У МЕРЕЖАХ ЗВ'ЯЗКУ <i>Безрук В.М., Гальченко К.Р.</i>	83
5	КЕРУВАННЯ РЕСУРСАМИ ДЛЯ ВІРТУАЛІЗОВАНИХ МЕРЕЖЕВИХ ФУНКЦІЙ <i>Скулиш М.А., Суліма С.В.</i>	97
6	ОПТИМІЗАЦІЯ ОБЧИСЛЕНЬ У ДАТА-ЦЕНТРИ ЗА КРИТЕРІЄМ ЕНЕРГОЕФЕКТИВНОСТІ <i>Гвоздецька Н.А., Глоба Л.С., Прокопець В.А., Степурін О.В.</i>	127
7	ДИНАМІЧНИЙ РОЗПОДІЛ РЕСУРСІВ У ВІРТУАЛІЗОВАНИЙ МЕРЕЖІ МОБІЛЬНОГО ОПЕРАТОРА УКРАЇНИ З ПІДВИЩЕННЯМ ЕНЕРГОЕФЕКТИВНОСТІ ОБРОБКИ ДАНИХ <i>Гвоздецька Н.А., Глоба Л.С., Прокопець В.А.</i>	151
8	МОДИФІКАЦІЯ МЕТОДУ ВЕРТИКАЛЬНОГО ХЕНДОВЕРУ ДЛЯ ГЕТЕРОГЕННОЇ МЕРЕЖІ <i>Глоба Л.С., Кірюшкін Р.О., Курдеча В.В.</i>	177

ВСТУП

Інфокомунікації – сучасна концепція конвергентної взаємодії телекомунікацій та інформаційних технологій з метою забезпечення своєчасної якісної передачі та обробки інформації.

Із впровадженням у світі технологій 4G LTE та розвитком тенденції переходу до технологій 5G постає питання збільшення швидкості та ефективності передачі та обробки даних. Все більшого поширення набувають сервіси Internet of Things (Інтернет речей) та сервіси міжмашинної взаємодії (M2M). Ці сервіси продукують великі масиви даних у режимі реального часу, що тягне за собою появу концепції BigData. Поряд із цим у світі гостро стоїть проблема економії електроенергії, що наразі споживається центрами обробки даних у величезних обсягах.

Такі виклики сучасного світу інфокомунікацій спричиняють необхідність пошуку нових наукоємних технологій оптимізації та управління при вирішенні задач планування та проектування інфокомунікаційних мереж. В монографії викладені деякі сучасні результати досліджень ряду авторів з питань оптимізації та управління в інфокомунікаційних мережах з урахуванням сукупності техніко-економічних вимог, включаючи швидкість передачі та обробки даних, енергоефективність, оптимальне використання ресурсів та інше. Це визначає актуальність опублікування даної колективної монографії, в якій викладені деякі результати досліджень авторів по вирішенню задач оптимізації та управління в інфокомунікаційних мережах з використанням сучасного математичного апарату та високоефективних обчислювальних засобів. Кожен розділ монографії являє собою завершене викладення результатів досліджень, отриманих авторами цього розділу.

Матеріали 1 розділу підготував Безрук В.М. У цьому розділі розглядаються деякі особливості методів багатокритеріальної оптимізації, що використовуються при розв'язанні різних типів задач оптимізації проектних рішень у мережах зв'язку з урахуванням сукупності показників якості. Наведено методи вирішення багатокритеріальних задач оптимізації при застосуванні безумовного критерію переваги, що приводить до отримання підмножини Парето-оптимальних рішень на множині допустимих варіантів. Для звуження підмножини Парето до єдиного переважного проектного варіанту застосовується умовний критерій переваги. Розглянуто різні методи формування умовного критерію переваги з урахуванням додаткової інформації від експертів, що дають можливість звуження підмножини Парето до єдиного переважного рішення. У наступних розділах 2 – 4 наводяться деякі приклади, що ілюструють практичні особливості використання розглянутих методів при багатокритеріальному аналізі і виборі оптимальних проектних варіантів для різних засобів зв'язку.

Матеріали 2 розділу підготували Чеботарьова Д.В., Безрук В.М. У цьому розділі розглянуто особливості застосування методів багатокритеріальної

оптимізації на етапі номінального планування радіомережі системи стільникового мобільного зв'язку (СМЗ) 2-го та 3-го покоління. Вирішена задача оптимізації радіомережі СМЗ як задача дискретного вибору Парето-оптимальних проектних варіантів при врахуванні сукупності показників якості. Розглянуто також особливості застосування методів багатокритеріальної оптимізації при номінальному плануванні транспортної мережі СМЗ. Наведені результати вибору оптимальної топології транспортної мережі СМЗ з урахуванням сукупності показників якості.

Матеріали 3 розділу підготували Безрук В.М., Скорик Ю.В., Власова В.О. У даному розділі розглянуто особливості вибору переважного проектного варіанту з урахуванням сукупності показників якості при застосуванні умовного критерія переваги, що заснований на використанні методу аналізу ієрархій (МАІ). Цей метод полягає в отриманні суджень експертів по парним порівнянь різних елементів проблеми вибору. В результаті обробки цих даних за строго формалізованими процедурами обчислюється вектор глобальних пріоритетів, компоненти якого визначають пріоритетність вибору переважного варіанту системи. Досліджено практичні особливості застосування МАІ для вибору переважного варіанту на прикладі різних типів засобів зв'язку, зокрема, мовних кодеків для ІР телефонії, системи масового обслуговування заявок у вузлах комутації мережі, проектних варіантів систем мобільного зв'язку 3-го покоління, технологій мобільного зв'язку 4-го покоління, типів мобільних телефонів, протоколів маршрутизації в ad-hoc мережах, модуляції в цифрових системах зв'язку, систем телевізійного мовлення стандартів DVB-T і ATSC, протоколів маршрутизації в бездротовій сенсорно-актуаторній мережі.

Матеріали 4 розділу підготували Безрук В.М., Гальченко К.Р. У даному розділі розглядаються особливості математичного формулювання та вирішення задач оптимальної маршрутизації у мережах зв'язку. Оскільки, практичні умови роботи мереж зв'язку вимагають врахування декількох показників якості (метрик) при виборі оптимальних маршрутів, тому розглянуто багатокритеріальний підхід до вирішення задач оптимальної маршрутизації у мережах зв'язку. Наведено метод формування підмножини Парето-оптимальних варіантів маршрутизації, що дають можливість організації багатошляхової маршрутизації у мережі зв'язку. Наводиться приклад вирішення задачі оптимальної маршрутизації у мультисервісній мережі зв'язку з урахуванням сукупності показників якості.

Матеріали 5 розділу підготували Скулиш М.А., Суліма С.В. У даному розділі запропоновано алгоритм оптимального розподілу ресурсів у віртуалізованих мережах. Метою алгоритму надання ресурсів є виділення достатньої їх кількості для мережевих функцій, так що їх SLA можна задовольнити навіть у присутності пікового навантаження. В основі будь-якого алгоритму надання ресурсів лежать два питання: скільки ресурсів надавати і коли. Ці питання глибоко досліджуються авторами розділу.

Матеріали 6 розділу підготували Глоба Л.С., Степурін О.В., Гвоздецька Н.А., Прокопеч В.А. У даному розділі запропоновано підхід до підвищення енергоефективності обчислень для серверного кластера як інформаційно-телекомунікаційної одиниці інфраструктури ЦОД, який відрізняється одночасним врахуванням параметрів енергоефективності та продуктивності при розподілі задач. Розроблено алгоритм підвищення енергоефективності обчислень на основі запропонованого підходу, який складається з етапу попередньої індивідуальної атестації серверного кластера та етапу енергоефективного розподілу задач. Врахування індивідуальних залежностей енергоспоживання серверів від їх завантаженості є основною відмінною рисою запропонованого підходу.

Матеріали 7 розділу підготували Глоба Л.С., Прокопеч В.А., Гвоздецька Н.А. У даному розділі проаналізовано ефективність методу динамічного розподілу ресурсів у віртуалізованій мережі мобільного оператора України з використанням механізмів підвищення енергоефективності обробки даних. Основною метою методу динамічного розподілу ресурсів є підтримка обсягу ресурсів, достатнього для виконання вимог SLA (Service Level Agreement – Угода про рівень обслуговування) при уникненні надмірності. Це дозволяє більш ефективно використовувати апаратні ресурси для економії електроенергії та зменшення викидів CO₂. Ефективність методу перевірено для віртуалізованого ядра EPC (Evolved Packet Core) мережі мобільного зв'язку України. Для додаткового підвищення енергоефективності обробки даних у дослідженні застосовано підхід PCPB (Power Consumption and Performance Balance), який запропоновано у розділі 6.

Матеріали 8 розділу підготували Кірюшкін Р.О., Курдеча В.В., Глоба Л.С. У даному розділі проведено огляд основних методів вертикального хендвера та зроблений висновок про неефективність існуючих рішень у бездротових гетерогенних мережах з багатокласовими груповими викликами мобільних вузлів. Описано модифікацію методу вертикального хендвера за рахунок зміни алгоритму прийняття рішень для різних типів викликів мобільного вузла що дозволило підвищити ефективність зв'язку у гетерогенних безпроводових мережах. Розглянуто створену математичну модель алгоритму прийняття рішення, який застосовує методи TOPSIS і MULTIMOORA. Проведено імітаційне моделювання методу прийняття рішення, який базується на запропонованій математичній моделі та доведено його ефективність у гетерогенних бездротових мережах.

5. КЕРУВАННЯ РЕСУРСАМИ ДЛЯ ВІРТУАЛІЗОВАНИХ МЕРЕЖЕВИХ ФУНКЦІЙ

Скулиш М.А., Суліма С.В.

5.1 Проблематика віртуалізації мережеских функцій

Надання послуг в телекомунікаційній галузі традиційно ґрунтується на тому, що мережескі оператори впроваджують фізичні пропріетарні пристрої та обладнання для кожної функції, яка є частиною певного сервісу. Крім того, сервісні компоненти мають чіткі ланцюги і/або порядок, які повинні бути відображені в топології мережі і в локалізації сервісних елементів. Це, в поєднанні з вимогами до високої якості, стабільності і строгим дотриманням протоколу, привело до тривалих циклів продукту, дуже низької гнучкості обслуговування і сильної залежності від спеціалізованих апаратних засобів.

Проте, вимоги користувачів до більш різноманітних і нових послуг з високими швидкостями передачі даних продовжують збільшуватися. Таким чином, телекомунікаційні сервіс-провайдери (Telecommunication Service Provider – TSP), повинні відповідним чином і постійно набувати, зберігати і експлуатувати нове фізичне обладнання. Це не тільки вимагає високих і швидко змінюваних навичок для техніків що експлуатують та управляють цим обладнанням, а й вимагає щільного розміщення мережевого обладнання, такого як базові станції. Все це призводить до високих капітальних і експлуатаційних витрат для TSP.

Як правило, оператори мобільного зв'язку справляються з підвищеними навантаженнями трафіку шляхом розширення/покращення загальної потужності інфраструктури відповідним чином. Проте, це все більш і більш важко реалізувати за рахунок збільшення капітальних/операційних витрат (CAPEX/OPEX) в світлі низької рентабельності інвестицій (ROI). Окрім низького ROI, надмірне резервування ресурсів більше не вважається життєздатною стратегією, щоб задовольнити збільшення трафіку, так як відповідно до [1], до 80% обчислювальної потужності базових станцій і до половини потужності ядра мережі є невикористаними. Це призводить до низького використання мережеских ресурсів, а також до високого рівня споживання енергії, що знижують економічну ефективність мережі для операторів мобільного зв'язку.

Принцип NFV спрямований на перетворення мережеских архітектур шляхом впровадження мережеских функцій в програмному забезпеченні, що може працювати на стандартній апаратній платформі. Крім того, він спрямований на перетворення традиційних мережеских операцій, оскільки програмне забезпечення може бути легко переміщене, або створено сутність в різних місцях без необхідності використовувати нове обладнання. NFV має багато переваг, від поліпшення операційної ефективності і зниження енергоспоживання до коротших інтервалів розгортання/оновлення і майже оптимального використання мережеских ресурсів, оскільки будівельні блоки можуть виділятися і перерозподілятися під час виконання в залежності від вимог.

З масовим розповсюдженням LTE оператори стикаються зі збільшенням у трафіку сигналізації. Оскільки абоненти стають дедалі активнішими в соціальних мережах і одночасно залучають все більше прикладних програм, смартфони можуть створювати величезні обсяги сигналізації, коли вони взаємодіють із мережею.

Проникнення мобільного широкосмугового зв'язку буде збільшуватися, і як наслідок, зростання трафіку сигналізації значно перевищить відповідний приріст трафіку даних. Компанія Nokia Siemens Networks [2] прогнозує, що зростання трафіку сигнальних повідомлень буде на 50% швидше, ніж зростання трафіку даних протягом найближчих кількох років. Проте збільшення використання смартфонів є лише одним із факторів вибуху сигнального трафіку.

Спонукальні чинники сигнального трафіку (рис. 5.1):

- оператори все більше покладаються на поінформованість про обслуговування, щоб забезпечити кращий досвід для користувачів смартфонів; у поєднанні з диференційованою тарифікацією, результат у збільшенні сигналізації становить до 30%;
- зростаюча кількість підключених до Інтернету мобільних пристроїв міжмашинної взаємодії та прикладних програм із високими вимогами мобільності призведе до значної сигналізації;
- Voice over LTE вимагає масштабованої пропускної здатності сигналізації та диференційованої якості сприйняття сервісів в режимі реального часу;
- виникаючі тенденції та нові бізнес-моделі;
- вимагають масштабованості площин користувача та управління;
- зростаюче використання онлайн-маркетингу, хмарного сховища, контент-бізнесу та прикладних програм;
- збільшення кількості клієнтів, які мають декілька персональних пристроїв;
- конвергенція таких галузей, як медіа, соціальні мережі, охорона здоров'я, енергетика та екологічні послуги тощо.

Еволюція мережі	Еволюція розумних пристроїв	Еволюція додатків і сервісів
Абоненти мереж LTE генерують утричі більше сигналізації у EPS	Двозначний ріст у впровадженні смартфонів в поєднанні з реалізаціями операційних систем конкретних постачальників призводить до збільшення сигналізації до 50%.	Сигналізація OverTheTop-гравців, включаючи багатий на відео контент, онлайн-ігри та мобільні додатки, знаходиться поза контролем операторів.

Рис. 5.1 Різноманітні причини зростання трафіку сигналізації [4]

Отже, NFV дозволяє операторам масштабувати послуги мережі з більшою деталізацією та оперативністю, ніж сьогодні, коли продуктивність мережевого сервісу визначена статично на найбільший прогнозований пік, що веде до надлишковості ресурсів. Для цього відповідні методи автоматизованого керування конфігурацією ресурсів гетерогенної мережі є надзвичайно важливими.

5.2. Опис методу відображення віртуальних вузлів на фізичні вузли

Підхід базується на спільному розташуванні індивідуальних ланцюгів сервісів ядра мережі на фізичній мережі, де ланцюг сервісів ядра мережі позначає послідовність мережевих функцій мобільного ядра, яку потік трафіку повинен пройти [9]. Припускаємо, що віртуальні мережеві функції мобільного ядра мають таку ж функціональність і інтерфейси як і мережеві елементи ядра архітектури 3GPP LTE EPC.

Фізична мережа задана у вигляді графа $SN = (N, L)$, де N є множиною фізичних вузлів і L – множиною каналів. Кожен канал $l = (n_1, n_2) \in L$, $n_1, n_2 \in N$ має максимальну пропускну здатність $c(n_1, n_2)$ і кожен вузол $n \in N$, пов'язаний з певними ресурсами c_n^i , $i \in R$, де R – множина типів ресурсів. Множина усіх точок агрегації трафіку (Traffic Aggregation Point – TAP), тобто кластерів eNodeB, в мережі позначається $T \subseteq N$. Для кожного вузла $n \in N$, $virt_n$ є бінарним параметром, який вказує, чи є вузол n віртуальним, $phys_n^j$ – бінарний параметр, який вказує, чи є вузол n виділеним апаратним блоком функції типу $j \in F$, де F – є множиною типів мережевих функцій.

Віртуальне ядро мобільної мережі представлено множиною ланцюгів сервісів (один ланцюг на TAP), які вбудовуються в фізичну мережу.

Вимоги смуги пропускання каналу між двома функціями, $j_1 \neq j_2$, $(j_1, j_2) \in E$, що відносяться до ланцюга сервісів TAP $t \in T$ позначається як $d_t^{(j_1, j_2)}$. $d_t^{j,i}$ – кількість ресурсу типу i , що виділяється для мережевої функції j TAP t . $s_t^{j,i}$ позначає час обробки запиту на ресурсі типу i мережевої функції j TAP t однією одиницею ресурсу для випадку віртуальної мережевої функції. Задане значення оброблених запитів в секунду виділеного апаратного функціонального блоку n типу j , і позначено як μ_n^j . Вимоги до запитів в секунду мережевої функції j , що відноситься до ланцюга сервісів TAP t , позначаються як M_t^j .

Метою оптимізації – є знаходження розташування ланцюгів сервісів ядра мережі (тобто розміщення мережевих функцій ядра мережі та розподіл ресурсів, а також визначення шляхів передачі трафіку між ними) таким чином, щоб мінімізувати витрати на зайняті ресурси каналів і вузлів у фізичній мережі, при цьому задовільняючи вимоги трафіку. Сформулюємо цільову функцію (формула (5.1)) у вигляді лінійної комбінації (з ваговими коефіцієнтами a , b , c) трьох вартісних виразів: базова вартість $cost(n)$, яка має місце, якщо будь-яка мережева функція розміщується на фізичному вузлі $n \in N$, вартість зайнятої одиниці ресурсів $cost(i, n)$ на фізичному вузлі n і вартість зайнятої одиниці пропускну здатності $cost(n_1, n_2)$ на фізичному каналі $(n_1, n_2) \in L$. Також прийемо, що L_t – максимальна допустима затримка для TAP $t \in T$, $L(n_1, n_2)$ – мережева затримка для каналу $(n_1, n_2) \in L$.

Вирази (5.1)–(5.9) представляють собою постановку оптимізаційної задачі нелінійного програмування. Змінні $x_n^{t,j}$ вказують, чи мережева функція j пов'язана з TAP $t \in T$ розташовується на фізичному вузлі n . Для $j = \text{TAP}$, $x_n^{t, \text{TAP}}$ – не змінні, а вхідні параметри, які вказують де TAP $t \in T$ знаходиться, тобто

$$x_n^{t, \text{TAP}} = \begin{cases} 1 & \text{якщо } t = n, \\ 0 & \text{інакше.} \end{cases}$$

Аналогічно, змінні $f_{(n_1, n_2)}^{t, (j_1, j_2)}$ вказують, чи фізичний канал $(n_1, n_2) \in L$ використовується для шляху між j_1 і j_2 для ТАР $t \in T$.

Таким чином, постановка задачі оптимізації може бути сформульована як описано далі.

Знайти:

$$\begin{aligned} \min_{x_n^{t,j}, f_{(n_1, n_2)}^{t, (j_1, j_2)}, d_t^{j,i}} & (a \cdot \sum_{n \in N} x_n \cdot cost(n) + b \cdot \\ & \sum_{n \in N} \sum_{t \in T} \sum_{j \in V} \sum_{i \in R} x_n^{t,j} \cdot d_t^{j,i} \cdot cost(i, n) + c \cdot \sum_{(n_1, n_2) \in L} cost(n_1, n_2) \cdot \\ & \sum_{t \in T} \sum_{(j_1, j_2) \in E} f_{(n_1, n_2)}^{t, (j_1, j_2)} \cdot d_t^{(j_1, j_2)}) \end{aligned} \quad (5.1)$$

За наявності обмежень:

$$\sum_{n \in N} x_n^{t,j} = 1 \quad \forall t \in T, j \in V \quad (5.2)$$

$$x_n^{t,j} \leq d_t^{j,i} \quad \forall t \in T, j \in V, n \in N, i \in R \quad (5.3)$$

$$\sum_{t \in T} \sum_{j \in V} x_n^{t,j} \cdot d_t^{j,i} \leq c_n^i \quad \forall n \in N, i \in R \quad (5.4)$$

$$\sum_{t \in T} \sum_{(j_1, j_2) \in E} f_{(n_1, n_2)}^{t, (j_1, j_2)} \cdot d_t^{(j_1, j_2)} \leq c(n_1, n_2) \quad \forall (n_1, n_2) \in L \quad (5.5)$$

$$\begin{aligned} \sum_{(n,w) \in L} f_{(w,n)}^{t, (j_1, j_2)} - f_{(n,w)}^{t, (j_1, j_2)} &= x_n^{t,j_1} - x_n^{t,j_2} \\ \forall t \in T, n \in N, (j_1, j_2) \in E \end{aligned} \quad (5.6)$$

$$x_n^{t,j}, f_{(n_1, n_2)}^{t, (j_1, j_2)} \in \{0,1\} \quad \forall t \in T, j \in V, n \in N, (j_1, j_2) \in E, (n_1, n_2) \in L \quad (5.7)$$

$$\mu_t^j \geq M_t^j \quad \forall t \in T, j \in V \quad (5.8)$$

$$\mu_t^j = \sum_{n \in N} \left(x_n^{t,j} \cdot \sum_{i \in R} \frac{d_t^{j,i}}{s_t^{j,i}} \cdot virt_n + x_n^{t,j} \cdot \mu c_n^j \cdot phys_n^j \right) \quad \forall t \in T, j \in V$$

$$\sum_{(j_1, j_2) \in E} \sum_{(n_1, n_2) \in L} f_{(n_1, n_2)}^{t, (j_1, j_2)} \cdot L(n_1, n_2) \leq L_t \quad \forall t \in T \quad (5.9)$$

Змінні x_n є булевими змінними, що позначають чи будь-яка мережева функція розміщується на вузлі $n \in N$. Наприклад $x_n=0$ означає, що немає жодної мережевої функції з будь-якого ланцюга сервісів, яка розміщувалася би в цьому вузлі.

Вираз (5.2) гарантує, що для кожного ТАР/ланцюга сервісів розміщується тільки одна мережева функція кожного типу. Це забезпечується таким чином, що сума змінних $x_n^{t,j}$ для розташування мережевої функції j , пов'язаної з ТАР t на фізичному вузлі n по всім вузлам, дорівнює одиниці.

Вираз (5.3) гарантує, що розміщення ресурсів здійснюється на фізичних вузлах, які було обрано для розташування відповідних мережевих функцій. Оскільки гарантується, що для випадку, коли булева змінна $x_n^{t,j}=1$, змінна $d_t^{j,i}$

буде відмінна від нуля для кількості ресурсу типу i , що виділено для мережевої функції j для TAP t .

Вирази (5.4) і (5.5) являють собою обмеження на ресурси фізичних вузлів і каналів, тобто забезпечують той факт, що кількість задіяних на вузлі ресурсів не перевищує кількості задіяних ресурсів. Слід зауважити, що канал між двома мережевими функціями відображається на шлях у фізичній мережі. Таким чином, його вимоги до полоси пропускання впливають не тільки на мережеві ресурси фізичних вузлів, де мережеві функції розміщуються, але й на мережеві ресурси проміжних вузлів, які лежать на шляху (вираз (5.5)).

Вираз (5.6) представляє собою обмеження щодо збереження потоку для всіх шляхів у фізичній мережі, тобто що вхідний потік на вузлі дорівнює вихідному потоку.

Вираз (5.7) гарантує, що змінні у задачі розміщення функції мережі та відображення шляху є булевими.

Для того, щоб урахувати у моделі продуктивність, обмеження на значення запитів в секунду, виражені у виразі (5.8), необхідні, щоб гарантувати, що загальна швидкість обслуговування мережевої функції j TAP t , яка позначається як μ_t^j , перевищує необхідну величину. Швидкість обслуговування може приймати або константне значення для випадку спеціалізованого виділеного функціонального блоку ($phys_n^j=1$), яке рівне μ_n^j , або змінне значення для випадку віртуалізованого функціонального блоку ($virt_n=1$), яке залежить від кількості виділених функціональному блоку ресурсів $d_t^{j,i}$ та часу обслуговування одиницею ресурсу $s_t^{j,i}$.

Щоб обмежити затримки на каналах, обмеження на затримку, показане в виразі (5.9), також додається, щоб обмежити загальну затримку на всьому шляху.

На рис. 5.2 представлено приклад системи розподілу мережевих запитів.

Предбачається вирішення задачі (5.1)-(5.9) в офлайн режимі на початковому етапі. Згідно з рішенням, кожному блоку мережевої функції виділяється певна кількість ресурсів, на основі грубої оцінки його потреби в ресурсах. Миттєві потреби різних мережевих функцій динамічно задовольняються під час виконання таким чином, щоб задовольнити гарантії передбачені для кожної мережевої функції.

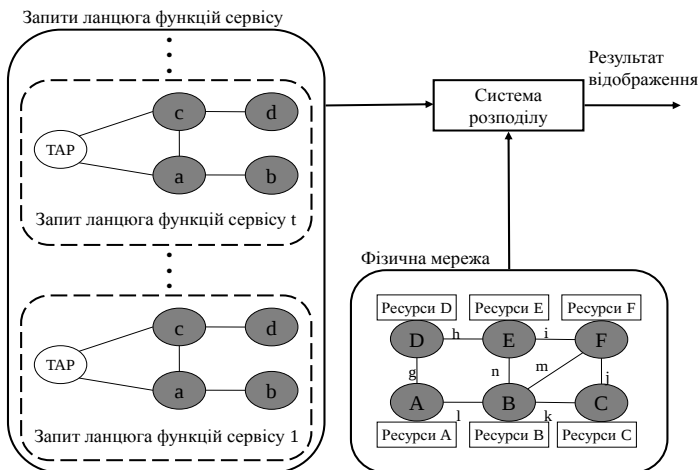


Рис. 5.2. Система виділення мережевих ресурсів – приклад топології

Увесь процес виділення ресурсів зображено на рис. 5.3.

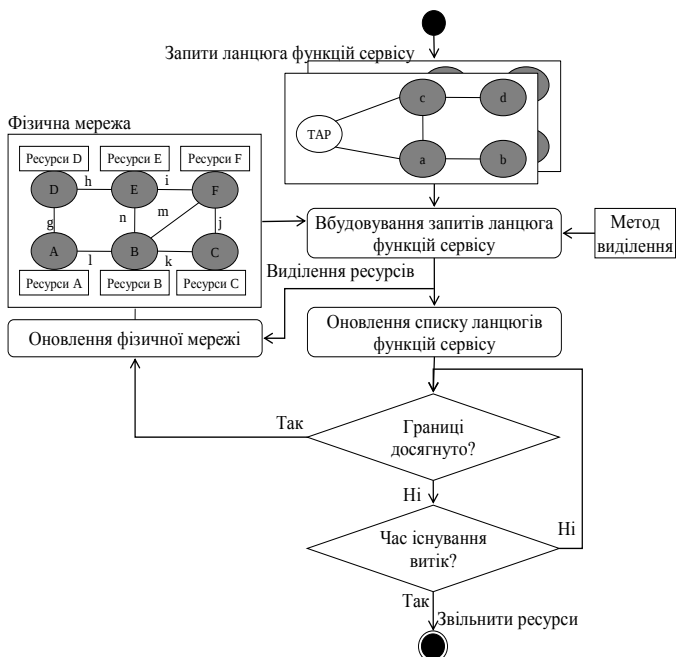


Рис. 5.3. Життєвий цикл мережевого запиту – діаграма активності

Процес виділення ресурсів починається з приходом запита ланцюга функцій сервісу, як показано на рис. 5.3. Метод розміщення та виділення ресурсів початкового етапу (тобто (5.1)-(5.9)) використовується для вбудовування запитів ланцюга функцій мережі, він приймає на вхід поточний стан фізичної мережі (наприклад, доступні ресурси CPU, пам'яті та пропускної здатності) і самі запити.

Метод направлений на мінімізацію використання ресурсів, при цьому забезпечуючи заданий рівень якості обслуговування. Запропонований метод дозволяє телекомунікаційному оператору мінімізувати капітальні та операційні витрати використовуючи парадигму хмарних обчислень та значно покращити якість сприйняття.

5.3 Опис задачі конфігурації активних ресурсів протягом дня

Метою алгоритму надання ресурсів є виділення достатньої їх кількості для мережевих функцій, таким чином що вимоги до якості обслуговування навантаження можна задовольнити протягом дня. В основі будь-якого алгоритму надання ресурсів лежать два питання: скільки надавати і коли [3].

Скільки надавати: Для вирішення питання про те, скільки ресурсів виділити для кожної мережевої функції, будемо аналітичну модель. Модель приймає в якості вхідних даних інтенсивність надходження вхідних запитів і вимоги обслуговування окремого запиту, і обчислює кількість ресурсів, необхідних мережеві функції, щоб впоратися з вимогами.

Коли надавати: Рішення про те, коли надавати ресурси, залежить від динаміки навантажень. Телекомунікаційні навантаження зазнають довгострокових змін, таких як вплив години дня або сезонні ефекти, а також короткострокових коливань таких як несподівані натовпи. У той час як довгострокові коливання можуть бути передбачені заздалегідь, спостерігаючи за змінами в минулому, короткострокові коливання менш передбачувані, а в деяких випадках, не передбачувані. Запропонована методика використовує два різних методи для роботи в умовах змін, які спостерігаються в різних часових масштабах. Використовується прогностичне керування ресурсами для оцінки навантаження і відповідного керування, а також реактивне керування ресурсами для виправлення помилок у довгострокових прогнозах або для реагування на непередбачені натовпи людей.

Розглянемо мережу, в якій функціонує кілька мережевих функцій [4]. Передбачається, що кожна така мережева функція вказує бажану вимогу до якості обслуговування (QoS); при цьому в даному випадку передбачаємо, що вимоги до QoS визначені в термінах цільового часу відповіді. Метою системи є забезпечення того, що середній час відповіді (або деякий процентиль часу відповіді), який спостерігається запитами мережевої функції, не перевищує бажаний цільовий час відповіді. Загалом, кожен вхідний запит обслуговується декількома апаратними та програмними ресурсами на сервері, такими як CPU, NIC, диск і т.д. Припускаємо, що заданий цільовий час відповіді розділяється на кілька значень часу відповіді для конкретних ресурсів по одному для кожного такого ресурсу. Таким чином, якщо кожен запит на кожному ресурсі не витрачає

часу більше, ніж призначене цільове значення, то загальний цільовий час відповіді для сервера буде задоволений. Завдання розділення зазначеного значення часу відповіді сервера на значення часу відповіді для конкретного ресурсу виходить за рамки даного дослідження; у роботі передбачається, що такі конкретні для ресурсу значення часу відповіді задані.

Для простоти викладу припускаємо, що в системі наявний лише один тип ресурсу.

Формально, L_i позначає цільовий час відповіді мережевої функції i і T_i – спостережуваний середній час відповіді, тоді мережеві функції потрібно виділити таку кількість ресурсів, щоб $T_i \leq L_i$.

Використаємо таку постановку задачі щоб одержати механізм динамічного виділення ресурсів, який описано далі.

5.4. Система динамічного керування ресурсами

Щоб виконати динамічне виділення ресурсів, засноване на вищевказаному формулюванні, на кожному сервері необхідно буде використовувати три компоненти: (1) модуль моніторингу, який вимірює навантаження і показники продуктивності кожної мережевої функції (такі як інтенсивність надходження запитів, середній час відповіді тощо), (2) модуль прогнозування, який використовує вимірювання з модуля моніторингу для оцінки характеристик навантаження в найближчому майбутньому, і (3) модуль розподілення ресурсів, який використовує ці оцінки навантаження для визначення кількості ресурсів, яку необхідно виділити мережевим функціям. На рис. 5.4 показані ці три компоненти [4].

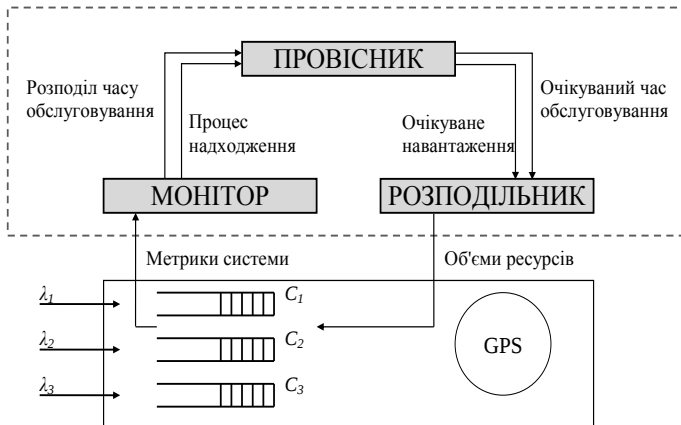


Рис. 5.4. Динамічне виділення ресурсів

Тепер представимо систему гібридного керування. Як вже відмічалось, навантаження часто мають передбачувані шаблони. Проте, можуть бути значні відхилення від цих моделей через пульсуючий характер навантажень телекомунікаційних систем. Надлишкові заявки або раптові сплески запитів можуть привести до погіршення якості обслуговування. Крім того, існують витрати та ризики, що пов'язані зі змінами у системі, і, таким чином, прагнуть уникнути частих змін виділення ресурсів. Ґрунтуючись на цих спостереженнях, пропонуємо гібридне рішення системи керування виділенням ресурсів, яке поєднує в собі прогностичну складову з реактивною. Ідея підходу полягає в тому, що фіксуються періодичні і стійкі моделі навантаження на основі історичних даних, які називаємо базовим навантаженням, а потім проактивно (і прогностично) керуємо ресурсами мережі. Під час виконання, різниця між фактичним і передбаченим навантаженням обробляється реактивним способом. Розроблена система також враховує витрати і ризики, пов'язані з керуванням ресурсами.

На рис. 5.2 показана концептуальна архітектура пропонованого рішення:

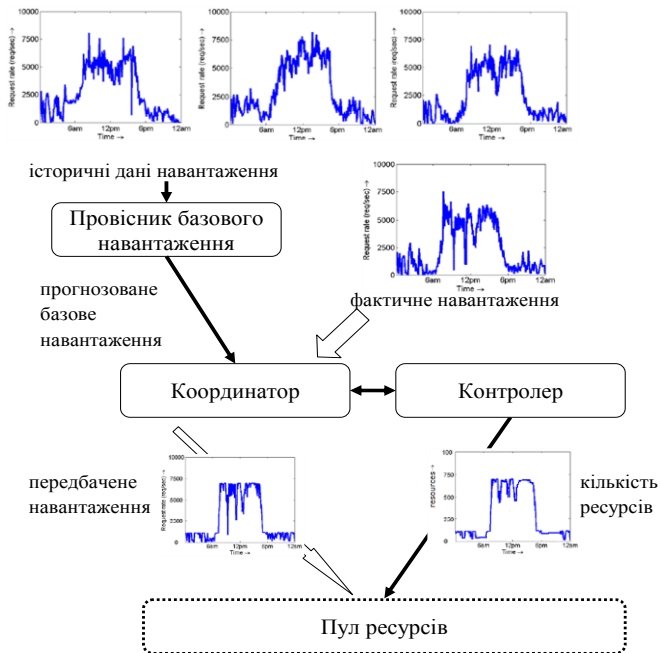


Рис. 5.5. Система гібридного керування

1. Провісник базового навантаження аналізує історичні дані навантаження і визначає закономірності, які формують базове навантаження.

2. Координатор направляє запити навантаження на сервери, а також обмінюється даними з контролером для надання інформації про вхідне навантаження.
3. Контролер оцінює і виділяє відповідну кількість ресурсів необхідну для обробки базового навантаження і надлишкових запитів. Наприклад, по даному навантаженню (змінювана у часі інтенсивність надходження запитів на рис. 5.5), він генерує розподіл ресурсів (кількість ресурсів на рис. 5.5).

Далі, спочатку представимо огляд модуля моніторингу та керування, який відповідає за виконання онлайн-вимірів. Після чого представимо опис у часовій області моделі масового обслуговування для ресурсів. Нарешті, представимо методи прогнозування, що використовуються для динамічної оцінки параметрів для цієї моделі.

5.5. Онлайн моніторинг та керування

Модуль онлайн моніторингу відповідає за вимірювання різних метрик системи і мережових функцій. Ці показники використовуються для оцінки параметрів моделі системи і характеристик навантаження.

Динамічний «інтелектуальний» моніторинг на мережі реалізується встановленням головної прикладної програми моніторингу на координатор та відповідних підконтрольних їй агентів на кожен мережовий блок [5].

Прикладна програма, встановлена на координаторі, здійснює аналіз моніторингу, а підконтрольні їй агенти здійснюють збір інформації, її модифікацію та відправлення на координатор. Диференціація інтенсивності моніторингу окремих сегментів мережі дає можливість усунути надлишкові процеси моніторингу та обробки даних, забезпечуючи необхідну інтенсивність лише на тих вузлах, де це необхідно. Використовуючи такий підхід можливо підвищити гнучкість керування мережевими елементами, і, як наслідок, ефективніше використовувати ресурси пристроїв рівня керування. Крім того, володіючи актуальною службовою інформацією, можна здійснювати оптимальне керування навантаженням у мережі. Це дозволить значно зменшити витрати на електроенергію та збільшити час «життя» елементів мережі.

Використовуючи традиційні методи керування станом ресурсів мережі, надлишкова службова інформація значно збільшується, що може негативно впливати на ефективність роботи мережі загалом через завантаженість каналів. Тому з метою забезпечення процесів керування системою, що побудована на основі принципів технології NFV, потрібно здійснити модернізацію механізмів їх роботи. Пропонується застосовувати метод динамічного керування конфігурацією ресурсів.

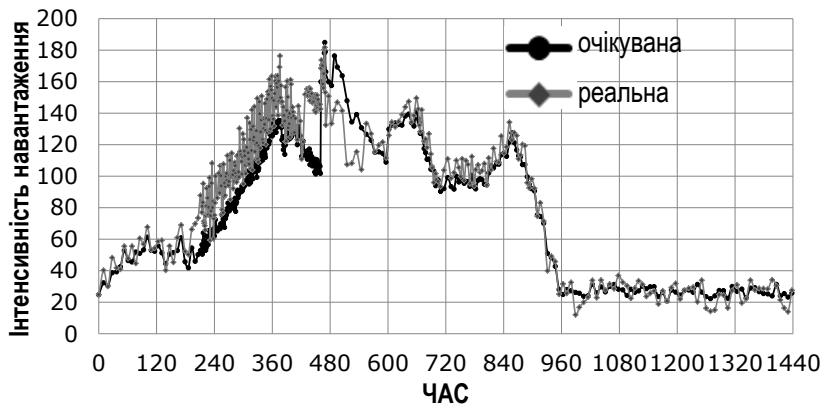
Процес ґрунтується на динамічній зміні інтервалу часу сталої конфігурації ресурсів залежно від відхилення фактичного значення навантаження від прогнозованого на основі зібраних у режимі реального часу даних стану мережевої функції. Якщо це значення збільшується, то скорочується інтервал сталої конфігурації.

Формула (5.22) описує принцип зміни частоти керування:

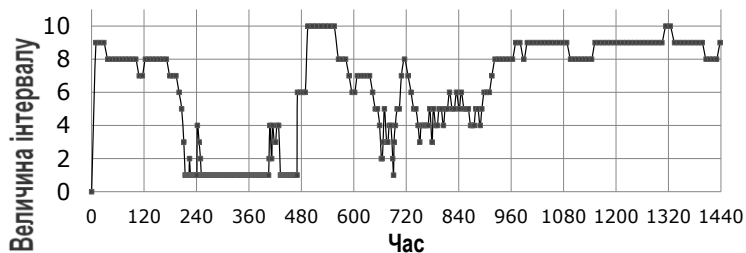
$$W(t) = \left(W_{base} - K \cdot \sum_{j=t-h}^{t-1} \frac{(\lambda_{obs}(j) - \lambda_{basepred}(j))^+}{\max \lambda_{basepred} - \lambda_{basepred}(j) + 1} \cdot W_{base}; W_{min base} \right), \quad (5.22)$$

де W – інтервал сталої конфігурації, W_{base} – базове значення інтервалу, $W_{min base}$ – мінімальна допустима величина базового інтервалу, K – коефіцієнт зміни тривалості сталої конфігурації, $\lambda_{obs}(t)$ – реальна інтенсивність надходження навантаження протягом інтервалу t , $\lambda_{basepred}(t)$ – середньостатистична передбачена інтенсивність надходження заявок під час інтервалу t , $\max \lambda_{basepred}$ – максимальна базова передбачена інтенсивність надходження заявок, h – кількість минулих спостережень, які враховуються.

На рис. 5.6 показано результати адаптації до відхилення реального навантаження від передбаченого на мережевій функції, отримані в процесі моделювання в системі Mathcad.



(а) Навантаження мережевої функції



(б) Зміна інтервалу сталої конфігурації відповідно до відхилення реального навантаження від очікуваного

Рис. 5.6. Зміна частоти адаптації

Вікно адаптації (W): вікно адаптації є часовим інтервалом між двома послідовними запусками алгоритму адаптації. Таким чином, минулі вимірювання використовуються для прогнозування навантаження для

наступного інтервалу часу W , і система адаптується для цього інтервалу часу. Як можна буде побачити далі, запропонований опис моделі масового обслуговування у часовій області розглядає період часу рівний вікну адаптації для оцінки середнього часу відповіді T_i мережевої функції, і ця модель оновлюється після інтервалу W або коли досягається заданий поріг.

Далі, визначаємо і дискретизуємо шаблони у прогнозованому навантаженні. Мета полягає в тому, щоб представити денний шаблон в навантаженні дискретизуючи його запити у послідовні, непересічні часові інтервали з єдиним репрезентативним значенням вимоги в кожному інтервалі. Запропоновано алгоритм для знаходження невеликої кількості часових інтервалів, таких, що відхилення від фактичної вимоги зведено до мінімуму. Підтримувати невелику кількість інтервалів важливо, так як більше число інтервалів означає більш часті зміни виділення ресурсів і, таким чином, більш високі ризик і витрати. Визначимо формально дискретизацію [4].

Дискретизація навантаження: Маючи часовий ряд X на області $[v, \tau]$, часовий ряд Y на тій же самій області є характеристизацією/дискретизацією навантаження X якщо $[v, \tau]$ може бути розділена на m послідовних непересічних часових інтервалів, $\{[v, \tau_1], [\tau_1, \tau_2], \dots, [\tau_{m-1}, \tau]\}$, так що $X(j)=r_i$, для всіх j у i^{mu} інтервалі, $[\tau_{i-1}, \tau_i]$.

Зверніть увагу, за визначенням, будь-який часовий ряд, X , є дискретизацією сам по собі. Встановлюємо $v=0$ і $\tau=0$ період навантаження. У подальшій дискусії припускаємо період у 24 години, на основі аналізу періодичності.

Ідея дискретизації подвійна. По-перше, потрібно точно відобразити вимоги. Для досягнення цієї мети, репрезентативні значення, r_i , для кожного інтервалу, $[\tau_{i-1}, \tau_i]$, мають бути якомога ближчими до фактичних значень часового ряду в інтервалі $[\tau_{i-1}, \tau_i]$. По-друге, виділення ІТ-ресурсів не відбувається безоплатно [7]. З цієї причини, потрібно уникати занадто великої кількості інтервалів і, отже, занадто великої кількості змін в системі, так як це не практично і може призвести до багатьох проблем (наприклад, втрати продуктивності, зносу серверів, нестабільності системи і т.д.) , Таким чином, слід мінімізувати помилку, внесену дискретизацією, і кількість інтервалів в дискретизації. Пропонується рішення для дискретизації часових рядів таким чином, щоб зменшити обидві величини.

$$\sum_{i=1}^m [\sum_{\tau=\tau_{i-1}}^{\tau_i} u(r_i - X(\tau))] + f(m) \rightarrow \min \quad (5.23)$$

Вираз (5.23) є цільовою функцією яку хочемо мінімізувати, де X являє собою часовий ряд і $f(m)$ є функцією вартості кількості змін або інтервалів, m . Метою виразу (5.23) є одночасна мінімізація помилки репрезентації навантаження і кількості змін. У деяких випадках, можна було б віддати перевагу, мінімізації квадрату кількості змін (або будь-якій іншій функції кількості змін). Вибір кращої функції вартості повинен здійснюватися мережевими адміністраторами з урахуванням конкретних потреб мережі під керуванням [8]. Для даної роботи встановлюємо $f(m)=l*m$, де l є константою нормалізації. Цільова функція виражає мету, мінімізуючи нормалізовану

кількість змін і помилку репрезентації. Функція вартості репрезентативної помилки використовується для кількісної оцінки похибки подання навантаження. Зверніть увагу, що в найзагальнішому випадку, як кількість змін, так і репрезентативна помилка можуть бути сформульовані як функції корисності.

Функція вартості кількості змін, наприклад, може відображати вартість електроенергії, а також грошову вартість пов'язану зі зносом на перезавантаження, яка базується приблизно на вартості заміни диска і номінальному часі напрацювання на відмову [9].

Оптимальну величину інтервалу визначаємо задаючи різні значення кількості інтервалів – обчислюємо значення виразу (5.23) і обираємо найкраще, тобто мінімальне, при цьому на кожному інтервалі: $r_i = \max_{\tau \in (\tau_{i-1}, \tau_i)} X(\tau)$.

Приклади репрезентації значень часового ряду представлені на рис. 5.5, де помилка репрезентації для випадку інтервалів у 10 хвилин складає 7%, а для випадку 60 хвилин – вже 19%.

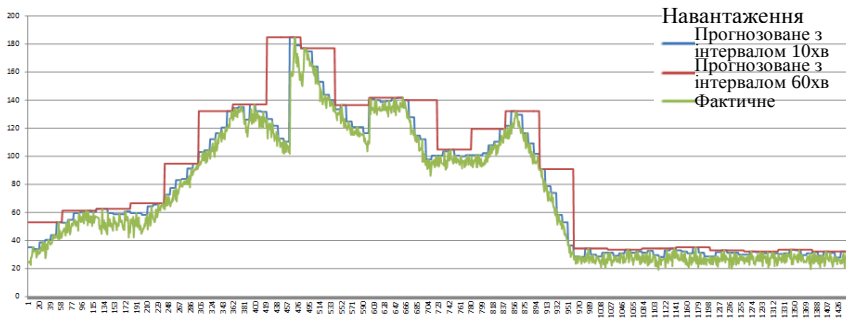


Рис. 5.7. Репрезентація значень навантаження в залежності від різної величини інтервалу адаптації

На рис. 5.8 показано основні етапи процедури визначення базового інтервалу, описаної вище. Псевдокод розглянутого алгоритму представлено на рис. 5.9.



Рис. 5.8. Схема визначення базового інтервалу

$interval \leftarrow maxinterval$

$i_{interval} \leftarrow 0$

$startpoint \leftarrow beginning$

while $interval > 0$

$startpoint \leftarrow 0$

$endpoint \leftarrow startpoint + interval$

while $startpoint < ending$

$maxvalue \leftarrow \max[Load(startpoint) \dots Load(endpoint)]$

for $j \in startpoint \dots endpoint$

$BaseLoad_{interval}(j) \leftarrow maxvalue$

$DiffLoad_{interval}(j) \leftarrow BaseLoad_{interval}(j) - Load(j)$

$startpoint \leftarrow endpoint$

$endpoint \leftarrow startpoint + interval$

$i_{interval} \leftarrow i_{interval} + 1$

$interval \leftarrow interval - 1$

for $j \in 1 \dots maxinterval$

$Minfunc_j \leftarrow \sum_k DiffLoad_j(k) + I * i_j$

$MinInt \leftarrow argmin(Minfunc)$

$MinCost \leftarrow min(Minfunc)$

Рис. 5.9. Псевдокод схеми визначення базового інтервалу

З іншого боку, додатково можемо ввести умову, що різниця між сусідніми репрезентативними значеннями не повинна бути меншою певного заданого порогу. Тобто для всіх інтервалів проводиться аналіз r_i , якщо $|r_i - r_{i+1}| < \varepsilon$,

де ε – деякий поріг, тоді значення i та $(i+1)$ об'єднуються, а нове репрезентативне значення обирається як $\max(r_i, r_{i+1})$.

5.6. Розподілення ресурсів серед мережевих функцій

Модуль розподілу ресурсів викликається періодично (кожне вікно адаптації або при досягненні порогу) для динамічного розділення ресурсного об'єму між різними мережевими функціями, які працюють на загальних серверах в мережі. Для охоплення перехідної поведінки навантажень мережевих функцій, спочатку представимо опис у часовій області моделі масового обслуговування ресурсів. Ця модель використовується для визначення потреби в ресурсах мережевої функції на основі очікуваного нею навантаження і цільового часу відповіді та заснована на підході, запропонованому у [4].

5.6.1. Модель масового обслуговування у часовій області

Як описано вище, алгоритм адаптації запускається кожні W часові одиниці. Нехай q_i^0 позначає довжину черги на початку вікна адаптації. Нехай λ_i позначає оцінку інтенсивності надходження заявок і μ_i позначає оцінку інтенсивності обслуговування у наступному вікні адаптації (тобто, на наступні W часові одиниці). Як ці значення оцінюються покажемо пізніше. Тоді, припускаючи що значення λ_i і μ_i є константними, довжина черги в будь-який момент часу t всередині наступного вікна адаптації задається як

$$q_i(t) = [q_i^0 + (\lambda_i - \mu_i)t]^+, \quad (5.24)$$

де x^+ позначає $\max(0, x)$.

Вочевидь, що обсяг завдань в черзі в момент часу t є сумою початкової довжини черги і обсягу завдань, які прибувають в цьому інтервалі, мінус обсяг завдань, які обслужилися в цей період. Крім того, довжина черги не може бути негативною.

Оскільки ресурс моделюється як GPS сервер, інтенсивність обслуговування запиту мережевої функції дорівнює d_i/s_i , де d_i – це кількість ресурсів мережевої функції і s_i – це середній час обслуговування запиту однією одиницею ресурсу. Отже, інтенсивність обслуговування запиту

$$\mu_i = d_i/s_i, \quad (5.25)$$

Відповідно до виразу 5.24, середня довжина черги під час вікна адаптації визначається за формулою:

$$q_i = \frac{1}{W} \int_0^W q_i(t) dt \quad (5.26)$$

В залежності від конкретних значень q_i^0 , інтенсивності надходження λ_i і інтенсивності обслуговування μ_i , черга може стати порожньою одного або більше раз протягом вікна адаптації. Щоб включити лише непусті періоди черги

обчислюючи q_i , розглядаємо наступні сценарії, основуючись на припущенні про постійні значення μ_i і λ_i :

1. Зріст черги: Якщо $\mu_i < \lambda_i$, тоді черга мережевої функції зростатиме протягом вікна адаптації і черга буде непустою протягом всього вікна адаптації.
2. Зменшення черги: Якщо $\mu_i > \lambda_i$, тоді черга починає зменшуватись протягом вікна адаптації. Момент часу t_0 в який черга стає пустою визначається як $t_0 = q_i^0 / (\mu_i - \lambda_i)$. Якщо $t_0 < W$, то черга стає пустою всередині вікна адаптації, в іншому випадку черга продовжує зменшуватись але залишається непустою протягом вікна (і прогнозується стати пустою у наступному вікні).
3. Постійна довжина черги: Якщо $\mu_i = \lambda_i$, тоді довжина черги залишається фіксованою ($= q_i^0$) протягом всього вікна адаптації. Отже, непустий період черги дорівнює або 0, або W залежно від значення q_i^0 .

Позначимо тривалість всередині вікна адаптації для якої черга непуста як W_i . Тоді, вираз 5.26 можна переписати як

$$q_i = \frac{1}{W} \int_0^{W_i} q_i(t) dx \quad (5.27)$$

$$= \frac{W_i}{W} \left[q_i^0 + \frac{W_i}{2} (\lambda_i - \mu_i) \right] \quad (5.28)$$

Визначивши середню довжину черги протягом наступного інтервалу адаптації, отримуємо середній час відповіді T_i у той самий інтервал часу. Нас цікавить середній час відповіді в найближчому майбутньому. Інші показники, такі як довгостроковий середній час відповіді також можуть бути розглянуті. T_i оцінюється як сума середньої затримки черги і часу обслуговування запиту протягом наступного інтервалу адаптації. Таким чином,

$$T_i = \left(\frac{q_i + 1}{\mu_i} \right) \quad (5.29)$$

Підставляючи вираз 5.25 в цей вираз, отримуємо

$$T_i = \frac{(q_i + 1)}{\left(\frac{d_i}{s_i} \right)}, \quad (5.30)$$

де q_i визначається виразом 5.27. Значення q_i^0 , λ_i , s_i отримуються використовуючи методики вимірювання та передбачення описані далі.

Такий опис моделі в часовій області має наступну ключову характеристику: параметри моделі залежать від її поточних характеристик навантаження (λ_i, s_i) і поточного стану системи (q_i^0). Відповідно, ця модель застосовується в онлайн сценарії реагування на динамічні зміни в навантаженні, і не робить ніяких припущень стаціонарного стану.

Далі опишемо, як ця модель використовується в динамічному розподілі ресурсів.

5.6.2. Задача розподілу ресурсів

Як пояснювалося раніше, обсяг ресурсів, який виділяється для мережевої функції, залежить від заданого цільового часу відповіді і розрахункового навантаження. Мережевій функції потрібно виділити кількість ресурсів так, що $T_i \leq L_i$. Цей вираз можна переписати з використанням виразу (5.30), як

$$\left(\frac{s_i}{d_i}\right) \cdot (q_i + 1) \leq L_i, \quad (5.31)$$

звідки кількість ресурсів, виділена мережевій функції d_i повинна задовольняти умову

$$d_i \geq \left(\frac{s_i}{L_i}\right) \cdot (q_i + 1) \quad (5.32)$$

5.7. Метод прогнозування навантаження

Онлайн метод розподілу ресурсів описаний в попередньому пункті у значній мірі залежить від точної оцінки навантаження, що може надійти. У цьому пункті представляємо метод, що використовує минулі спостереження для оцінки майбутнього навантаження мережевої функції.

Навантаження, яке зазнає мережева функція, можна охарактеризувати двома доповнюючими розподілами: процесом надходження запитів і розподілом часу обслуговування. Разом ці розподіли дозволяють охопити інтенсивність навантаження і його мінливість. Запропонована методика вимірює різні параметри, що визначають ці розподіли протягом певного періоду часу, і використовує ці вимірювання для прогнозування навантаження для наступного вікна адаптації.

5.7.1. Оцінка інтенсивності надходження

Найважливішим параметрами, які характеризують процес надходження, є інтенсивність надходження запитів та її точна оцінка, що дозволяє моделі масового обслуговування у часовій області оцінювати середню довжину черги для наступного вікна адаптації.

Провісник навантаження заснований на методиці, запропонованій в [3], який використовує минулі спостереження за навантаженням, щоб передбачити пікову вимогу, яка буде зазнаватися протягом часу T . Для простоти викладу припустимо, що $T=1$ година. У цьому випадку, провісник оцінює пікову вимогу протягом наступної однієї години на початку кожної години. Щоб зробити це, він зберігає історію інтенсивності надходження сеансів, яка мала місце протягом кожної години дня протягом останніх декількох днів. Потім для кожної години генерується гістограма, використовуючи спостереження для цієї години за останні декілька днів (див. рис. 5.10). Кожна гістограма дає розподіл ймовірностей інтенсивності надходження за цю годину. Пікове навантаження

для певної години оцінюється як високий процентиль розподілу інтенсивності надходження за цю годину (див. рис. 5.10). Таким чином, використовуючи хвіст розподілу інтенсивності надходження для передбачення пікової вимоги, прогностичний метод надання ресурсів може виділити достатні ресурси, щоб впоратися з навантаженням найгіршого випадку, якщо воно надійде. Крім того, моніторинг вимог у кожен годину дня дозволяє провіснику охоплювати ефекти години дня.

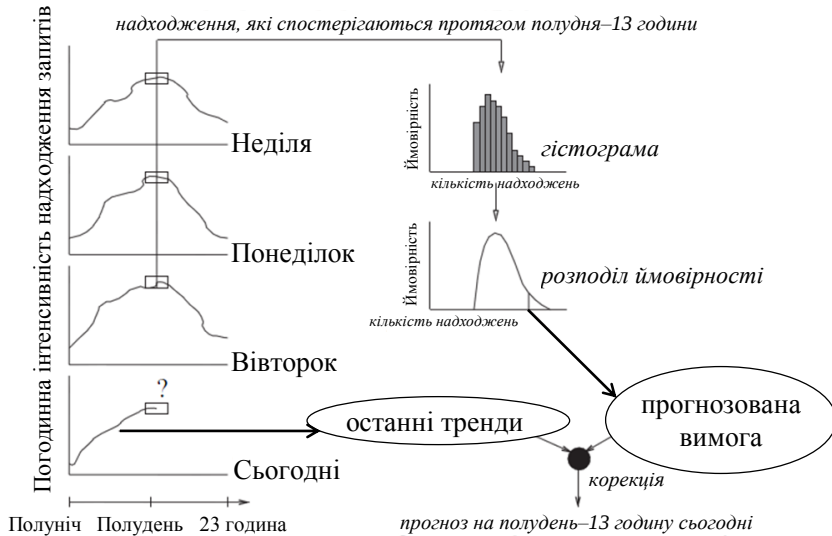


Рис. 5.10. Схема передбачення навантаження

На додаток до використання спостережень за попередні дні, навантаження в останні кілька годин поточного дня можуть бути використані для того, щоб далі підвищити точність прогнозування.

Припустимо, що $\lambda_{basepred}(t)$ – базова передбачена інтенсивність надходження під час певного інтервалу позначеного t . Далі, нехай $\lambda_{obs}(t)$ – реальна інтенсивність надходження протягом цього інтервалу. Помилка передбачення це $(\lambda_{obs}(t) - \lambda_{basepred}(t))$. У випадку помилки передбачення протягом останніх декількох інтервалів, прогнозоване значення протягом наступного інтервалу коригується з використанням помилки, що спостерігається:

$$\lambda_{pred}(t) = \lambda_{basepred}(t) + \alpha \sum_{j=t-h}^{t-1} (1 - \alpha)^{t-1-j} \cdot (\lambda_{obs}(j) - \lambda_{basepred}(j))^+, \quad (5.33)$$

де другий вираз позначає зважену помилку передбачення за останні h інтервалів, а x^+ позначає $\max(0, x)$.

Таким чином, пропонується здійснювати прогнозування навантаження на наступний інтервал часу шляхом урахування прогнозу на основі довгострокової статистики та коригуванням його за моделлю експоненційного згладжування (exponentially weighted moving average), де помилки більш нових

минулих періодів мають більший ваговий коефіцієнт. α (константа згладжування) – коефіцієнт, що характеризує швидкість зменшення ваг, та приймає значення від 0 до 1, чим значення цього параметра ближче до одиниці, тим більше при прогнозі враховується вплив останніх рівнів ряду. Параметри моделі встановлюються адміністратором мережі відповідно до експерименту.

За допомогою прогнозованої пікової інтенсивності надходження для кожної мережевої функції, метод прогностичного надання ресурсів використовує модель для визначення кількості ресурсів, які мають бути виділені для кожного компоненту.

5.7.2. Оцінка часу обслуговування

Двом мережевим функціям з подібними інтенсивностями надходження запитів, але різними інтенсивностями обслуговування (наприклад, через різні розміри пакетів і т.д.), потрібно буде виділити різні кількості ресурсів [4].

Для оцінки інтенсивності обслуговування мережевої функції модуль прогнозування обчислює розподіл ймовірностей часу обслуговування запитів. Такий розподіл представлено гістограмою часу обслуговування запиту. По завершенні кожного запиту, ця гістограма оновлюється з урахуванням часу обслуговування цього запиту. Розподіл використовується для визначення очікуваного часу обслуговування s_i для запитів в наступному вікні адаптації. s_i може обчислюватись як середнє, медіана, або процентиль розподілу, отриманого з гістограми. Пропонується використовувати середнє значення розподілу для представлення часу обслуговування запитів мережевої функції.

5.7.3. Вимірювання довжини черги

Останнім параметром, необхідним для моделі розподілу ресурсів є довжина черги кожної мережевої функції на початку кожного вікна адаптації. Оскільки нас цікавить лише миттєва довжина черги q_i^0 , а не середні значення, вимірювання цього параметра є тривіальним – модуль моніторингу записує кількість запитів, що очікують в кожній черзі мережевої функції на початку кожного вікна адаптації.

5.8. Механізми моніторингу

Важливий аспект платформи полягає в здатності моніторингу її різних складових елементів [10]. Цей механізм дозволяє платформі адаптувати і масштабувати мережеві функції у відповідності з різними метриками, де сутності прийняття рішень порівнюють отримані тригери з вимогами описів сервісів окремих елементів, які повідомляються ним. Платформа забезпечує можливість відділення аспектів моніторингу функціонування мережі мобільного зв'язку від тих, які відносяться до інфраструктури NFV, підтримуючи існуючі механізми керування та оркестровки в зв'язку з цим. Конкретно, інтерпретація подій моніторингу OSS/BSS і MCE (Management Component Entity) з фізичної LTE EPC може бути використана для запуску створення віртуальної EPC (vEPC)

в інфраструктурі NFV, або навіть просто створення одного елемента vEPC, або компонента. Таким чином, доступні механізми моніторингу в платформі надаються на різних рівнях:

- VNF: коли VNF вбудовує функції моніторингу і передає зібрану інформацію в Element Management (EM) або VNFManger (VNFM). Система дозволяє функціям VNF безпосередньо направляти дані сутностям OSS/BSS і MCE з подіями моніторингу, що дозволяє їм приймати загальномеревеві рішення щодо керування процесами та функціями, з урахуванням також віртуалізованих функцій;
- EM: EM може реалізувати рішення про дії, на основі даних, зібраних з VNF, і направити їх на VNFM, або до OSS/BSS і MCE для дій на всю мережу;
- OSS/BSS і MCE: коли функція моніторингу не реалізована в EM, або перетинає кілька EM. Це може бути як точкою детектування, так і рішення. OSS/BSS також отримує події моніторингу LTE EPC, що дозволяє запускати зміни в віртуалізованій мережі мобільного зв'язку або її елементів;
- Оркестровщик NFV (NFVO): коли робота віртуалізованого сервісу зазнає впливу на рівні мережевого сервісу (тобто складається з декількох взаємопов'язаних VNF з загальним кінцем), то NFVO може бути покращено можливістю спільно розглядати різні звіти про моніторинг окремих VNF, і діяти, коли функціональні потреби мережевого сервісу в цілому піддаються впливу, запускаючи VNFM діяти над необхідними VNF;
- VNFM: він приймає рішення масштабування, основуючись на подіях VNF, відповідно до параметрів масштабування і моніторингу в їх дескрипторах. Цей об'єкт може використовувати цю інформацію, а також діяти у відношенні необхідних дій інфраструктури NFV, отримуючи події від менеджера віртуалізованої інфраструктури (Virtualized Infrastructure Manager – VIM).

Таким чином, результат моніторингу в розглянутій платформі спрямований в бік загальної поведінки масштабування віртуальних ресурсів відповідно до продуктивності мережевого сервісу VNF, трьома шляхами: 1) автоматичне масштабування, де VNFM моніторить події і запускає операцію масштабування при виконанні певних умов, 2) масштабування на вимогу, коли VNF, або його EM, викликає операцію масштабування за допомогою явного запиту на VNFM, або 3) запит керування, який ініціюється відправником (OSS/BSS або MCE) у напрямку до VNFM через NFVO.

Для того, щоб здійснювати моніторинг роботи площини користувача і площини керування, S1-U і S5 є точками спостереження для моніторингу продуктивності трафіку даних в той час, як S1-MME і S10 є точками спостереження для моніторингу продуктивності навантаження сигналізації через необхідні операції і активність абонентів для користувальницького обладнання. Інформація моніторингу з означених опорних точок спостереження збирається на OSS/BSS. Інформація від сутностей компонентів vEPC збирається на кожному EM і доставляється зрештою OSS/BSS.

5.9. Метод реконфігурації мережі після перевантаження або збою

Фізична мережа задана у вигляді графа $SN=(N,NE)$, де N є множиною фізичних вузлів і NE – множиною каналів. Кожен канал $(n_1,n_2) \in NE$, $n_1,n_2 \in N$ має максимальну пропускну здатність $c(n_1,n_2)$ і мережеву затримку $L(n_1,n_2)$, а кожен вузол $n \in N$ пов'язаний з певними ресурсами c_n^i , $i \in R$, де R – множина типів ресурсів. Мережа зв'язку представлена множиною сервісів (або запитів віртуальної мережі) K , які вбудовуються в фізичну мережу. Запит віртуальної мережі k , $k \in K$, можна представити як зважений граф $G_k=(V_k,E_k)$, де V_k є множиною віртуальних вузлів, що містить h_k елементів і позначається як $V_k=(v_{k,1},v_{k,2},\dots,v_{k,h_k})$, де $v_{k,j}$ означає j -у мережеву функцію у ланцюзі функцій k E_k є множиною віртуальних каналів $e_k(v_{k,j},v_{k,g}) \in E_k$. Вимоги смуги пропускання каналу між двома функціями, j_1 і j_2 , що відносяться до сервісу $k \in K$ позначається як $d_k^{(j_1,j_2)}$, $d_k^{j,i}$ – кількість ресурсу типу i , що виділяється для мережевої функції j сервісу k . Булеві змінні $x_n^{k,j}$ вказують, чи мережева функція j , пов'язана з $k \in K$, розташовується на фізичному вузлі n , змінні $f_{(n_1,n_2)}^{k,(j_1,j_2)}$ визначають, чи фізичний канал (n_1,n_2) використовується у шляху між j_1 та j_2 для запиту k . L_k – максимальна затримка для запиту k . $costN(i,n)$ – вартість зайнятої одиниці ресурсу i на фізичному вузлі n , і $costL(n_1,n_2)$ – вартість зайнятої одиниці пропускну здатності на фізичному каналі $(n_1,n_2) \in NE$. $suit_n^{k,j}$ означає що функція j з запиту k може бути розміщена на вузлі n .

MN являє собою множину вузлів керування, де $MN \subseteq N$, які відповідають за функціонування пропонованого механізму реконфігурації після відмови або перевантаження. Кожен керуючий вузол пов'язаний з одним або декількома вузлами фізичної мережі і виконує кроки, необхідні для відновлення належного функціонування мережі. Процес призначення вузлів керування і критерії, які враховуються при виборі вузлів керування, будуть досліджені далі.

У запропонованому підході припускаємо, що відображення запитів віртуальної мережі вже виконано і вузли керування у фізичній мережі призначаються так само, як описано нижче.

Процес відображення віртуальної мережі відбувається в два етапи [12]: відображення вузлів ($M_N : V_k \rightarrow N$) і відображення каналів ($M_L : E_k \rightarrow NE$).

Проблема, що розглядається далі, полягає в тому, як перемістити віртуальні вузли, розміщені на вузлі з відмовою, з метою мінімізації витрат на відновлення вузла відмови і періоду переривання сервісу. У запропонованому підході підкреслюється, що будь-яке переміщення віртуального вузла повинно виконуватись локально і бути координованим тільки вузлами керування.

Розглянемо підхід щодо оптимального розміщення вузлів керування у мережах, заснованих на NFV.

Припускаємо, що вузли керування (далі – менеджери) можуть розміщуватись в точках, де оператор вже має сайт, тобто вузли керування розміщуються в вузлах N .

При заданій кількості менеджерів A існує скінченна множина з $\binom{|N|}{A}$ можливих розташувань, відповідно, задача розміщення менеджерів є задачею багатокритеріальної комбінаторної оптимізації. Метою задачі є знаходження таких розташувань менеджерів з множини можливих розташувань розміру A –

$Pl_A = \{Pl \in 2^{|N|} \mid |Pl| = A\}$, що є оптимальними відповідно до деякої цільової функції.

Метою оптимізації є визначення місця розташування кожного менеджера при заданій їх кількості A , так що мінімізується функція загальних витрат $Cost_A(\{p_n: n \in N\})$, де p_n – булева змінна, яка рівна одиниці, якщо менеджер розміщується в точці n . Задача оптимізації буде мати вигляд:

$$\min_{\{p_n: n \in N\}} U_A \quad (5.34)$$

при обмеженні $\sum_n p_n = A$

Основною метою оптимального розміщення менеджерів є мінімізація затримок між вузлами і менеджерами в мережі. Проте, розглядати тільки затримки не достатньо. Розміщення менеджерів повинно також враховувати певні обмеження стійкості.

На рис. 5.11 показані різні питання, які необхідно враховувати при оцінці стійкості розміщення. Нижче коротко пояснимо ці питання, і що потрібно, щоб бути стійким по відношенню до них. На рис. 5.11 показано нормалізовані затримки між вузлами та інтенсивність навантаження на вузлах.

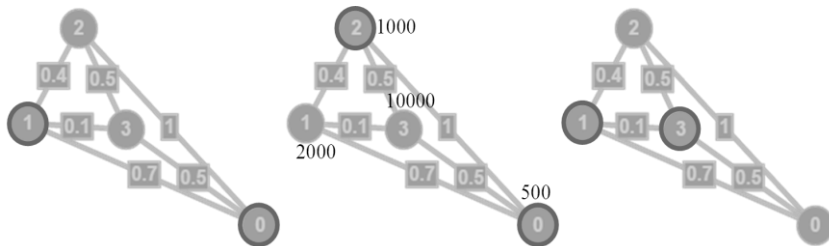


Рис. 5.11. Розміщення по різним критеріям: (а) мінімальної затримки до менеджера; (б) мінімального дисбалансу навантаження на менеджерів; (в) мінімальної затримки між менеджерами

Аналогічно до [13], припустимо, що вузли призначаються до їх найближчого менеджера, використовуючи в якості метрики затримки $dl_{g1,g2}$ між вузлом $g1$ і менеджером $g2$. Кількість вузлів на менеджера може бути незбалансованою – чим більше вузлів менеджер повинен контролювати, тим вище навантаження на цього менеджера. Якщо кількість запитів вузла до менеджера в мережі збільшується, аналогічно поводить себе і ймовірність додаткових затримок через черги в системі керування. Для того, щоб бути стійким від перевантаження менеджера, призначення вузлів різним менеджерам повинно бути добре збалансованим.

Цілком очевидно, що одного менеджера не досить, щоб досягти стійкості в мережі. Проте, коли кілька менеджерів розміщуються в мережі, логіка керування мережі розподіляється по декількох менеджерах і ці менеджери повинні синхронізуватися, щоб підтримувати несуперечний глобальний стан.

Залежно від частоти синхронізації між менеджерами, затримка між окремими менеджерами грає важливу роль.

1. Затримка вузол-менеджер

Аналогічно, як представлено в [14], на основі матриці dl , що містить відстані найкоротших шляхів між усіма вузлами, максимальний час затримки передачі між вузлом і менеджером для певного розміщення менеджерів може бути визначений як

$$U_A^{latency}(p) = \max ddc_n \quad (5.35)$$

ddc_n – максимальна затримка передачі від вузла мережі до менеджеру в точці n ;

ddc_n розраховується наступним чином:

$$ddc_n = \max_{g \in N} latency_g \cdot \pi_{g,n},$$

де $latency_g$ – затримка між менеджером та вузлом g , $latency_g = \min_{\{n: n \in N \cap p_n = 1\}} dl_{g,n}$;

$\pi_{g,n}$ – булева змінна, яка рівна одиниці, якщо вузол g обслуговується менеджером розміщеним в точці n .

Розглядаємо не середнє, але максимальне значення затримки, так як середнє приховує значення найгіршого випадку, які є важливими, коли стійкість повинна бути поліпшена.

2. Балансування навантаження менеджера

Залежно від ситуації, може бути бажаним наявність приблизно рівного навантаження на всіх менеджерах, аби жоден менеджер не перевантажувався, а інші не мали меншої роботи. Далі розглядаємо збалансований розподіл вузлів між менеджерами. Як формальну метрику вводимо баланс розміщення або, вірніше, дисбаланс, $U_A^{imbalance}$, тобто відхилення від повністю збалансованого розподілу, як різниця між навантаженням на найбільш завантаженому менеджері і найменш завантаженому менеджері.

$U_A^{imbalance}$ визначається в такий спосіб:

$$U_A^{imbalance}(p) = \max ldc_n - \min ldc_n, \text{ де } ldc_n > 0 \quad (5.36)$$

ldc_n – навантаження на менеджер в точці n ;

ldc_n розраховується наступним чином:

$$ldc_n = \sum_{g \in N} load_g \cdot \pi_{g,n}$$

де $load_g$ – коефіцієнт навантаження на вузол g .

3. Затримка менеджер-менеджер

Як останній аспект стійкого розміщення менеджерів, розглянемо як затримка між менеджерами може враховуватися при виборі розміщення менеджерів. Формально, затримка між менеджерами $U_A^{interlatency}$ визначається як

найбільша затримка між будь-якими двома менеджерами при заданому розміщенні:

$$U_A^{interlatency}(p) = \max_{\{n,g:n,g \in N \cap p_n=1, p_g=1\}} dl_{g,n}. \quad (5.37)$$

Загалом, без ілюстрації розміщень, розміщення з урахуванням затримки між менеджерами мають тенденцію розміщувати всіх менеджерів набагато ближче один до одного. Це збільшує максимальну затримку від вузлів до менеджерів.

4. Цільова функція оптимізації

Таким чином, застосовуємо адитивний критерій оптимізації мінімальної затримки до менеджера, мінімального дисбалансу навантаження на менеджерів, мінімальної затримки між менеджерами:

$$U_A = w_l \times U_A^{latency}(p) + w_i \times U_A^{imbalane}(p) + w_{il} \times U_A^{interlatency}(p), \quad (5.38)$$

де w_i – множина вагових коефіцієнтів.

На рис. 5.12 показані можливі рішення по розміщенню двох менеджерів у мережі з 10 вузлами. Оптимальне значення показує, що усі оптимізаційні цілі не можуть бути досягнуті водночас.

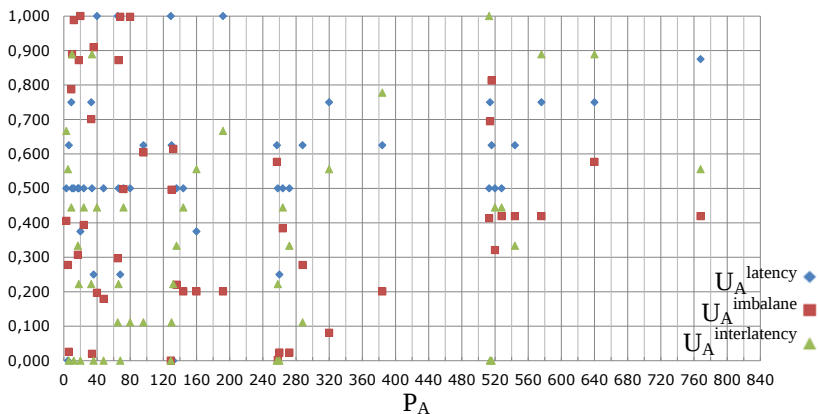


Рис. 5.12. Простір рішень оптимізаційної задачі розміщення менеджерів

5. Конфігурація вузлів менеджерів

Кожен менеджер (наприклад, вузол 3 на рис. 5.13) управляє вузлом 2) підтримує таблицю (наприклад, таблиця 5.1), що містить ідентифікатори віртуальних мереж, які мають віртуальні вузли розміщені на керованому фізичному вузлі [15]. Крім того, таблиця містить посилання на розміщені на фізичних вузлах віртуальні вузли (наприклад, вузли c та d на рис. 5.13) і суміжні до них віртуальні вузли (наприклад, вузли a та b на рис. 5.13). Ця таблиця

безперервно оновлюється (наприклад, кожен раз, коли прийнято або відхилено запит віртуальної мережі).

На рис. 5.13 показаний приклад фізичної мережі з двома вже вбудованими віртуальними мережами. Вузол з відмовою (вузол 2) був призначений менеджеру (вузлу 3), який відповідає за переміщення віртуальних вузлів.

Таблиця 5.1 Стани на менеджері 3

ID віртуальної мережі	Віртуальний вузол	Ємність	Прилеглі вузли
1	c	10	a на вузлі 0 b на вузлі 1
2	d	15	f на вузлі 4 e на вузлі 5

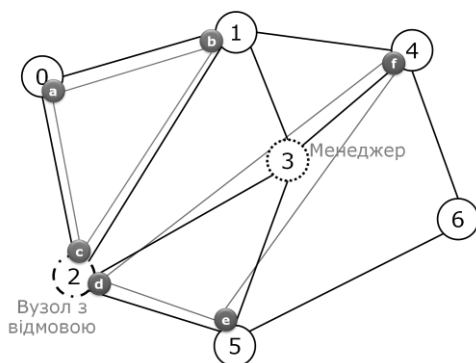


Рис. 5.13. Фізична мережа з вузлом з відмовою, на якому розміщується два віртуальних вузла

Алгоритм відновлення мережі після відмови може бути запропонованим виходячи з наступних міркувань.

Процес переміщення вузлів віртуальної мережі, розміщених на вузлі, який відмовив, v_{kj}^{fail} (де k – ідентифікатор віртуальної мережі, j – віртуального вузла), запускається, коли система відправляє запит на відновлення відповідному вузлу-менеджеру (наприклад, вузлу 3 на рис. 5.13). Процес відновлення для кожного порушеного вузла віртуальної мережі протікає в такий спосіб: менеджер направляє запит на відновлення до всіх вузлів фізичної мережі, на яких розміщаються віртуальні вузли, суміжні з ураженими віртуальними вузлами (наприклад, вузол 0 і вузол 1 на рис. 5.13). Кожен з цих вузлів будує дерево найкоротших шляхів (Shortest Path Tree – SPT) до всіх вузлів фізичної мережі на відстані не більше l (поріг встановлюється провайдером послуг), де коренем SPT виступає сам цей вузол. Менеджер використовує ці шляхи, щоб вибрати вузол з оптимальною відстанню до всіх вузлів фізичної мережі, де розташовані вузли віртуальної мережі прилеглі до несправного вузла. Цей вузол в кінцевому рахунку стає оптимальним кандидатом для розміщення віртуального

вузла з відмовою. Крім того, ємність кінцевих вузлів шляхів з SPT повинна бути не менше ємності віртуального вузла, розміщеного на вузлі з відмовою (наприклад, вузла с). Обираємо вузол з мінімальною вартістю шляху до всіх кореневих вузлів у деревах SPT та мінімальною вартістю обчислень.

На рис. 5.14 показано основні кроки алгоритму, які відпрацьовують запит на відновлення мережевої функції. Псевдокод алгоритму відновлення вузла після відмови описано на рис. 5.15 і він виконується для всіх $\{v_{kj} : x_n^{kj}=1 \ \& \ n = failed\}$.

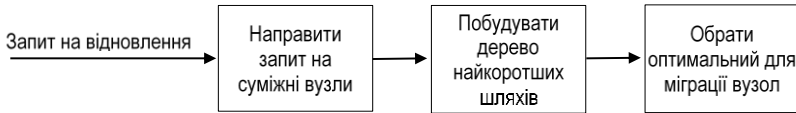


Рис. 5.14. Відновлення вузла з відмовою

У мережі існує ймовірність відмови вузла через перевантаження. Для виконання відновлення при відмові вузла в перевантаженій мережі, виконується процедура реконфігурації для міграції віртуальних вузлів, розміщених на перевантаженому фізичному вузлі.

Процес відновлення (рис. 5.16) починається з сортування всіх віртуальних вузлів, розміщених на перевантаженому фізичному вузлі. Критерієм сортування цих вузлів віртуальної мережі є ємність віртуальних вузлів. Потім виконується процедура відновлення на першому відсортованому вузлі віртуальної мережі, що має ємність, яка дорівнює перевантаженню, для переміщення на новий вузол фізичної мережі.

Коли навантаження або ресурси змінюються, деякі VNF можливо доведеться перемістити. Існує ймовірність того, що знайти новий вузол-кандидат для вузла віртуальної мережі, розміщеного на вузлі з відмовою, не вийде. В такому випадку виконується процедура реконфігурації для міграції одного або декількох віртуальних вузлів для переміщення розміщених вузлів віртуальної мережі. Розглянемо завдання міграції як завдання оптимізації, що спрямоване на мінімізацію загальної вартості міграції при обмеженнях допустимої затримки і обчислювальних ресурсів.

Метою оптимізації є знаходження розташування ланцюгів сервісів мережі (тобто розміщення мережевих функцій та розподіл ресурсів, а також визначення шляхів передачі трафіку між ними) таким чином, щоб мінімізувати витрати на зайняті ресурси каналів і вузлів у фізичній мережі, при цьому задовольняючи вимоги трафіку. Сформулюємо цільову функцію у вигляді лінійної комбінації (з ваговими коефіцієнтами a, b, c, e) вартісних виразів. Визначимо бінарну змінну $x_n^{kj} \in \{0,1\}$, для позначення того, що VNF j пов'язана з ланцюгом сервісу k розміщується на вузлі n після міграції. Індикатор $x_n^{kj}=0$ означає, що VNF j не розміщується на вузлі n після міграції; в іншому випадку j розміщується на вузлі n після міграції.

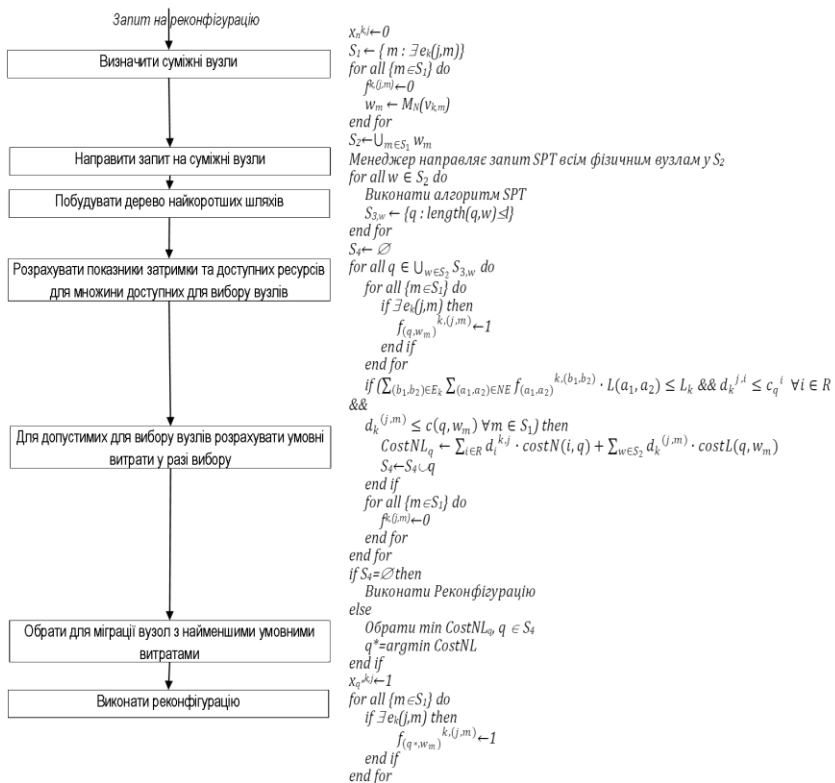


Рис. 5.15. Псевдокод алгоритму відновлення вузла після відмови



Рис. 5.16. Відновлення вузла з перевантаженням

Введемо бінарну змінну $y_n^{k,j}$ для індикації стану мережі перед міграцією. Вона схожа зі змінною $x_n^{k,j}$, $y_n^{k,j}=0$ означає, що VNF j сервісу k не перебуває на вузлі n до міграції; в іншому випадку, j розташована на вузлі n до міграції.

Таким чином, можемо використовувати індикатор $I^{k,j}$, щоб вказати чи VNF j сервісу k було переміщено в поточному процесі міграції.

$$I^{k,j} = \sum_{n \in N} x_n^{k,j} \cdot y_n^{k,j}$$

$I^{kj} = 0$ вказує, що VNF було переміщено в поточному процесі міграції, і $I^{kj} = 1$ вказує, що VNF не було переміщено.

x_n позначає чи n -ий фізичний сервер працює або ні після міграції.

$$x_n = \begin{cases} 1 & (\text{сервер } n \text{ працює}) \\ 0 & (\text{в іншому випадку}) \end{cases}$$

y_n позначає чи n -ий фізичний сервер працює або ні перед міграцією.

$$y_n = \begin{cases} 1 & (\text{сервер } n \text{ працює}) \\ 0 & (\text{в іншому випадку}) \end{cases}$$

Для того щоб розглянути ресурси, які споживаються при міграції та запуску серверів, вводимо такі вирази:

B_n позначає необхідні витрати b_n для запуску n -ого серверу.

$$B_n = b_n x_n (x_n - y_n);$$

$L_i^{k,j}(n \rightarrow n')$ позначає використання ресурсу i для міграції VNF j з ланцюгу сервісу k з серверу n на сервер n' .

$$L_i^{k,j}(n \rightarrow n') = l_i(d^{k,j}) + l'_i(d^{k,j}),$$

де $l_i(x)$ – функція використання ресурсу i для міграції з серверу, $l'_i(x)$ – функція використання ресурсу i для міграції на сервер.

Цільова функція буде визначатися як:

$$\begin{aligned} MCost = & a \cdot \sum_{n \in N} (B_n + x_n \cdot cost(n)) + b \cdot \sum_{n \in N} \sum_{k \in K} \sum_{j \in V} \sum_{i \in R} x_n^{k,j} \cdot \\ & d_i^{k,j} \cdot costN(i, n) + c \cdot \\ & \sum_{(n_1, n_2) \in NE} costL(n_1, n_2) \cdot \sum_{k \in K} \sum_{(j_1, j_2) \in E} f_{(n_1, n_2)}^{k, (j_1, j_2)} \cdot d_k^{(j_1, j_2)} + e \cdot \\ & \sum_{n \in N} \sum_{n' \in N} L_i^{k,j}(n \rightarrow n') x_{n'} (x_{n'} - y_n). \end{aligned} \quad (5.39)$$

Використовуючи наведені вище міркування, формулюємо задачу наступним чином.

Знайти: $\min MCost$

З обмеженнями:

$$\sum_{n \in N} x_n^{k,j} = 1 \quad \forall k \in K, j \in V \quad (5.40)$$

$$x_n^{k,j} \leq suit_n^{k,j} \quad \forall k \in K, j \in V, n \in N \quad (5.41)$$

$$\begin{aligned} \sum_{k \in K} \sum_{j \in V} x_n^{k,j} \cdot d_i^{k,j} + y_n^{k,j} \cdot (1 - I^{k,j}) \cdot l_i(d^{k,j}) + x_n^{k,j} \cdot (1 - I^{k,j}) \cdot \\ l'_i(d^{k,j}) \leq c_n^i \quad \forall n \in N, i \in R \end{aligned} \quad (5.42)$$

$$\sum_{t \in K} \sum_{(j_1, j_2) \in E} f_{(n_1, n_2)}^{k, (j_1, j_2)} \cdot d_k^{(j_1, j_2)} \leq c(n_1, n_2) \quad \forall (n_1, n_2) \in NE \quad (5.43)$$

$$\begin{aligned} \sum_{(n, w) \in L} f_{(w, n)}^{k, (j_1, j_2)} - f_{(n, w)}^{k, (j_1, j_2)} = x_n^{k, j_1} - x_n^{k, j_2} \quad \forall k \in K, n \in \\ N, (j_1, j_2) \in E \end{aligned} \quad (5.44)$$

$$\begin{aligned} x_n^{k,j}, f_{(n_1, n_2)}^{k, (j_1, j_2)} \in \{0, 1\} \quad \forall k \in K, j \in V, n \in N, (j_1, j_2) \in \\ E, (n_1, n_2) \in NE \end{aligned} \quad (5.45)$$

$$\sum_{(j_1, j_2) \in E} \sum_{(n_1, n_2) \in NE} f_{(n_1, n_2)}^{k, (j_1, j_2)} \cdot L(n_1, n_2) \leq L_k \quad \forall k \in K \quad (5.46)$$

$$\sum_{n \in N} x_n^{k, j} \sum_{i \in R} \left(\frac{1}{\frac{d_{k, j, i}}{s_{n, i}^{k, j}} - \lambda^{k, j}} \right) \leq P_k^j \quad \forall k \in K, j \in V \quad (5.47)$$

Отже, цільова функція (5.39) представляє собою лінійну комбінацію чотирьох вартісних виразів та направлена на мінімізацію: вартості запуску та використання серверу, використання ресурсів серверу, використання каналів зв'язку та використання ресурсів для міграції. Обмеження (5.40) гарантує однократність розміщення мережеских функцій, а (5.41) – адміністративну можливість розміщення на вузлі. Вирази (5.42) та (5.43) являють собою обмеження на ресурси фізичних вузлів і каналів, тобто забезпечують той факт, що кількість задіяних на вузлі ресурсів не перевищує кількості доступних. Вираз (5.44) представляє собою обмеження щодо збереження потоку для всіх шляхів у фізичній мережі, тобто що вхідний потік на вузлі дорівнює вихідному потоку. Вираз (5.45) гарантує, що змінні у задачі є булевими. Вирази (5.46) та (5.47) представляють собою обмеження на час передачі телекомунікаційними каналами та час обробки вузлами обслуговування відповідно, і забезпечують дотримання заданих часових вимог до обслуговування сервісу.

Висновки

1. Запропоновано модель організації ресурсів у дата центрах оператора мобільного зв'язку, яка дозволяє описати процес виділення ресурсів мережевими функціями у мережі оператора мобільного зв'язку.

2. Запропоновано метод визначення інтервалу часу сталої конфігурації ресурсів з урахуванням очікуваного навантаження, який дозволяє забезпечити ефективне керування сервісами у віртуалізованому середовищі.

3. Удосконалено розподілений метод локальної реконфігурації ресурсів, який розподіляє та перерозподіляє віртуальні мережеві функції в нормальному та аварійному режимі із забезпеченням ефективного використання ресурсів.

Література

1. Liquid Core [Online]. – Available at: <http://networks.nokia.com/portfolio/liquidnet/liquidcore>.

2. Signaling is growing 50% faster than data traffic [Online]. – Available at: <http://docplayer.net/6278117-Signaling-is-growing-50-faster-than-data-traffic.html>.

3. Uргаonkar B. Agile dynamic provisioning of multi-tier Internet applications / B. Uргаonkar, P. Shenoy, A. Chandra, P. Goyal, T. Wood // ACM Transactions on Autonomous and Adaptive Systems (TAAS). – 2008. – Vol. 3, No. 1. – P.1-39.

4. Chandra A. Dynamic resource allocation for shared data centers using online measurements / A. Chandra, W. Gong, P. Shenoy // 11th international conference on Quality of service. – Berkeley, CA, USA, 2003. – P. 381-398.
5. Коваль Б. В. Дослідження площини керування програмно-керованих мереж на основі розподіленої системи функцій віртуалізації / Б. В. Коваль, М. О. Селюченко, Г. В. Мельник, А. В. Ковальчук // Вісник Національного університету «Львівська політехніка». Радіоелектроніка та телекомунікації. – 2014. - № 796. - С. 164-175.
6. Gandhi A. Minimizing Data Center SLA Violations and Power Consumption via Hybrid Resource Provisioning / A. Gandhi, Y. Chen, D. Gmach, M. Arlitt, and M. Marwah // Green Computing : International Conference and Workshops, Orlando, FL, 25–28 July 2011 : proceedings. – IEEE, 2011. – P. 1–8.
7. Emmerson B. M2M: The internet of 50 billion devices / B. Emmerson // WinWin Magazine. – 2010. – P. 19-22.
8. Bianco A. Cost and performance trade-offs in reconfiguration strategies for WDM networks / A. Bianco, J. Finochietto, C. Piglione// 4th International Telecommunication Networking Workshop on QoS in Multiservice IP Networks. – Venice, Italy, 2008. – P. 95-100.
9. Chen Y. Managing Server Energy and Operational Costs in Hosting Centers / Y. Chen, A. Das, W. Qin, A. Sivasubramaniam, Q. Wang, N. Gautam // SIGMETRICS '05 Proceedings of the 2005 ACM SIGMETRICS international conference on Measurement and modeling of computer systems. – Banff, Alberta, Canada, 2005. – P. 303-314.
10. Jeon S. Virtualised EPC for on-demand mobile traffic offloading in 5G environments / S. Jeon, D. Corujo, R. L. Aguiar // 2015 IEEE Conference on Standards for Communications and Networking (CSCN). – Tokyo, Japan, 2016. – P. 275-281.
11. Baumgartner A. Mobile core network virtualization: A model for combined virtual core network function placement and topology optimization / A. Baumgartner, V.S. Reddy, T. Bauschert // 2015 1st IEEE Conference on Network Softwarization (NetSoft). – London, 2015. – P. 1-9.
12. Chowdhury N.M.M.K. Virtual Network Embedding with Coordinated Node and Link Mapping / N.M.M.K. Chowdhury, M. Rahman and R. Boutaba // INFCOM. – Rio de Janeiro, Brazil, 2009. – P. 1-9.
13. Hock D. Pareto-optimal resilient controller placement in SDN-based core networks / D. Hock, M. Hartmann, S. Gebert, M. Jarschel et al. // 25th International Teletraffic Congress (ITC). – Shanghai, 2013. – P. 1-9.
14. Heller B. The controller placement problem / B. Heller, R. Sherwood, N. McKeown // ACM HotSDN. – Helsinki, Finland, 2012. – P. 1-6.
15. Abid H. A novel scheme for node failure recovery in virtualized networks / H. Abid; N. Samaan // 2013 IFIP/IEEE International Symposium on Integrated Network Management (IM 2013). – Ghent, Belgium, 2013. – P. 1154-1160.